

Material Imprimible

Curso Creación de prompts

Módulo Inteligencia Artificial y ChatGPT

Contenidos:

- Contexto del panorama actual de la Inteligencia Artificial, su impacto global y la trascendencia de su integración en nuestras vidas y profesiones
- Introducción a la Inteligencia Artificial. Principales modelos de Inteligencia Artificial. Introducción a los Modelos de texto
- ChatGPT: qué es, cómo funciona, para qué sirve
- Open-AI: creación de una cuenta, navegación en la plataforma, y descubrimiento de sus funciones y comandos básicos
- Cuándo es conveniente usar ChatGPT y cuándo no
- Errores a evitar al usar ChatGPT

Contexto del panorama actual de la Inteligencia Artificial, su impacto global y la trascendencia de su integración en nuestras vidas y profesiones

La Inteligencia Artificial, también llamada I.A, lleva mucho tiempo desarrollándose y es el cúmulo de diferentes tecnologías que acaban de explotar.

Podemos decir que la misma tiene sus primeras apariciones en la historia de la humanidad muchísimo antes del siglo pasado, pero realmente es en 1950 cuando el matemático y lógico británico Alan Turing empezó a popularizarla hablando de ello con naturalidad.

Él en su ensayo *“Computing Machinery and Intelligence”* propuso la idea de que las máquinas podían pensar, introduciendo el famoso “Test de Turing” para determinar si una máquina puede exhibir un comportamiento inteligente indistinguible del de un ser humano.

Los años pasaron, se siguió investigando, y en el año 2000 aproximadamente empezamos a notar un incremento muy grande de Inteligencia Artificial, y el motivo principal es la gran capacidad de cómputo que tenemos.

En 2010 salieron teléfonos móviles que tenían la misma capacidad de computación que computadoras que teníamos años atrás. Ahora mismo es una barbaridad, cualquier computadora personal es muchísimo mejor que cualquiera del año 2000 y similar. Por lo tanto, todos tenemos una capacidad de cómputo muy grande, con mucha potencia.

De hecho, en la década del 2010 sale “Siri”, de Apple, que es el primer asistente virtual ampliamente utilizado en dispositivos móviles, capaz de realizar tareas mediante comandos de voz.

La Inteligencia Artificial ha avanzado mucho por el Hardware, porque ahora tenemos esa capacidad de poder llevar a la práctica aquello que años atrás era puramente teórico.

En concreto, nosotros en este curso vamos a hablar sobre texto, sobre la herramienta que vamos a ver. El texto al final es un área de la Inteligencia Artificial que trabaja en NLP, es decir, Procesamiento de Lenguaje Natural, que es una forma de trabajar diferente a la que se trabaja, por ejemplo, con modelos predictivos matemáticos y demás.

La generación de texto, los modelos generativos, se basan en un corpus de información que están formados por un montón de palabras y frases acumuladas en un mismo lugar, y entonces entrenamos modelos de Inteligencia Artificial con esa información.

En el año 2015 es fundada Open-AI, la empresa que crea ChatGPT. Esta es una empresa que se crea como un laboratorio de Inteligencia Artificial en el que sin ser una empresa con fines de lucro, busca recibir financiación para poder tener a los mejores científicos de Inteligencia Artificial desarrollando modelos, simplemente por el hecho de avanzar en este campo.

Un poco más tarde, en 2017, Google decidió liberar una de sus tecnologías. Aquí hacemos un paréntesis para contarles que la Inteligencia Artificial en aquella época no era un negocio. Quienes trabajaban o hacían cosas con Inteligencia Artificial eran pocos, no era lo normal; era una tecnología emergente que por primera vez en la historia te la podías instalar en tu computadora y empezar a desarrollar tus propios modelos de Inteligencia Artificial, pero no era algo tan *mainstream* como es hoy en día.

Por lo dicho, lo que hacían las grandes empresas cuando tenían un nuevo hallazgo, una nueva tecnología, era liberarla con el ánimo de que otras empresas pudiesen tomar esa tecnología, desarrollar algo con ella, y volver a liberar. Ese es un poco el pensamiento de la parte del Software Open Source, en el que todo el mundo intenta contribuir para que el pensamiento colectivo haga desarrollarse esa tecnología.

Entonces los investigadores de Google, en 2017, deciden compartir una tecnología que han desarrollado, que se llaman los Transformers. En este paper, que se hizo muy famoso en la época y es un hito en la computación, los investigadores explican que han encontrado una nueva forma de tratar el texto, y que a partir de ahora ya no es necesario entender el texto de forma lineal, sino que se lo puede entender por diferentes vías, por así decirlo.

Es decir, simplemente para que lo entendamos todos, ahora son capaces de comprender muchísimo más texto de lo que hacía antes. Por tanto, la Inteligencia Artificial, de la noche a la mañana, con esta nueva técnica, con esta nueva tecnología que desarrolla Google, es mucho más inteligente si hablamos en concreto de texto.

En el año 2019, Open-AI, este laboratorio de ideas, empieza a hablar de los Transformers. Dado que Open-AI fue fundada como un laboratorio de ideas, en su página web tienen todos los avances que han hecho durante todos los años, por lo que ustedes podrán

ingresar a ver todos los artículos. En este que vemos en pantalla, es donde empiezan a contar que están utilizando los transformes de Google, esa tecnología que Google había liberado, para generar sus propios modelos de Inteligencia Artificial generativa de texto.

Un poco más tarde, de repente, Open-AI publica este artículo que vemos en pantalla, en el que le cuenta al mundo que ha creado un modelo de Inteligencia Artificial llamado GPT, que lo que hace es generar texto basado en un histórico muy grande de información que le han dado al modelo, y que es el mayor generador de texto que la humanidad ha visto.

Es tan grande este modelo generativo de texto que en este mismo artículo, y luego en muchos otros, deciden no liberarlo. Y ahí es un poco donde empiezan a generar esa emoción entre las personas que estaban en ese mundo en aquel momento. Tan bueno es, que sería tan dañino liberarlo a día de hoy.

Entonces, en ese momento muchas personas pusieron a Open-AI en el mapa y comenzaron a fijarse en su GPT para ver si es presunción o realmente tienen un modelo tan potente como dice basado en toda la tecnología de Google.

Dos años más tarde, repentinamente, Open-AI dice que va a liberar GPT, que va a hacer un acceso a la API, y que cualquier persona, sin ningún tipo de restricción, podrá acceder al modelo.

En 2021, quienes tenían acceso a GPT eran personas que sabían programar, personas que sabían conectarse a una API, utilizar directamente códigos de programación. Hoy en día sigue siendo algo relativamente complejo acceder a este tipo de datos, es decir, no todo el mundo puede hacerlo, no todo el mundo puede trabajar con esa tecnología.

Un poco más tarde, en 2022, ahora sí deciden hacerlo totalmente público y aquí es donde la locura explota. Pero... ¿por qué explota? porque es la primera vez en la historia en la que personas, usuarios sin ningún tipo de conocimiento de programación, pueden comunicarse de forma directa con la Inteligencia Artificial generativa de texto seguramente más avanzada de la historia.

Después de todos estos años, ChatGPT convierte a la Inteligencia Artificial en algo usual, en algo en el que cualquier persona puede conectarse, puede hablar, puede mandar un mensaje sin necesidad de conocimientos técnicos. Da igual del sector que venga la

persona y los conocimientos que tenga, puesto que cualquier individuo puede utilizarlo, dejando de ser una cuestión de programadores, científicos de datos, etc.

Es a partir de entonces donde explota esa emoción, ese entusiasmo por la Inteligencia Artificial, ya que todas las personas que ahora están accediendo a esta tecnología se dan cuenta del potencial que tiene; potencial que ya habían visto años atrás programadores y científicos de datos, pero dado que era un sector muy pequeño, no se generaba revuelo. No obstante, a partir de este momento en el que esto es tan relevante, empieza a generar la agitación en la que estamos hoy en día.

Después de esto han pasado muchas cosas, pero claramente este es el momento crucial, el momento donde todo el mundo es capaz de acceder a ChatGPT y puede ver la capacidad que éste tiene de ayudarnos.

Finalmente diremos que aunque no lo parezca, lo que podemos hacer con ChatGPT sigue siendo algo técnico, porque para sacar el máximo partido de ChatGPT hay que saber cómo comunicarse con la Inteligencia Artificial, hay que entender cuál es el modelo que hay detrás; pero en este curso lo que vamos a hacer es precisamente eso.

La Inteligencia Artificial y sus principales modelos

En primer lugar vamos a definir a la **Inteligencia Artificial** como un campo de la informática que se centra en la creación de sistemas que pueden realizar tareas que normalmente requieren ser realizadas por humanos, como el aprendizaje, la toma de decisiones, y el reconocimiento de patrones.

Esta se basa en la idea de que un ordenador pueda ser programado para simular la forma en la que funciona el cerebro humano y así lograr realizar tareas que de otra manera requerirían la intervención de dicho ser.

La Inteligencia Artificial se divide en varias ramas: el aprendizaje automático, el PLN, que es el procesamiento del lenguaje natural, la robótica, etc. Nosotros nos vamos a centrar en el PLN.

Ahora veremos juntos qué es un **modelo**. Podemos decir que es una representación matemática o computacional que se utiliza para tomar decisiones o hacer predicciones según los datos de entrada que recibe.

Es decir, es como si fuera una caja negra, en la que sucede lo siguiente: en primer lugar, recibe datos de entrada, lo que es un *input*; a continuación, realiza un procesamiento de esos datos de entrada mediante métodos de aprendizaje automático y algoritmos, que son secuencias de instrucciones lógicas y matemáticas; y finalmente, el modelo produce una salida, un *output*, que es lo que utilizamos los humanos para tomar decisiones o realizar predicciones.

Ahora bien. Podemos mencionar los siguientes principales modelos de Inteligencia Artificial:

- Los modelos de texto, que son modelos basados en el PLN. Estos son capaces de realizar tareas relacionadas con la lectura y escritura, como hacer resúmenes y mapas mentales, traducciones, etc. Uno de los modelos más conocidos son los modelos GPT
- También tenemos modelos conversacionales o de diálogo, que son modelos de generación de texto, es decir, que están preparados para poder conversar, como es el caso, por ejemplo, de ChatGPT
- Otro modelo es el de generación de imágenes, que son una clase de modelos de aprendizaje profundo capaces de comprender visualmente conceptos. A partir de datos de entrada, este tipo de modelos es capaz de generar imágenes nuevas que no existían. El más conocido es el modelo GAN. Por ejemplo, una herramienta muy popular es Midjourney, que es un modelo de generación de imágenes. Nosotros le aportamos un *input* de texto y el modelo va a producir un *output* tipo imagen.
- Luego tenemos los modelos de generación de video, que son una clase de modelos de aprendizaje profundo capaces de generar secuencias de video a partir de datos de entrada. Uno de los más conocidos es el VGAN. En este caso nosotros le aportamos al modelo, por ejemplo, un *input* de texto, y el modelo lo que va a hacer es producir un *output* tipo video
- Finalmente encontramos los modelos de generación de audio, que son una clase de modelos de aprendizaje profundo capaces de generar nuevos sonidos y pistas de audio a partir de datos de entrada. Uno de los modelos más conocidos es el WaveNet. En este caso vamos a poder introducir un *input* de texto al modelo, y este va a producir un *output*, que va a ser una pista de audio

Dado que ya hicimos una introducción general a la Inteligencia artificial y a los diferentes modelos, ahora vamos a realizar una introducción a los modelos de texto, y para eso lo primero que tenemos que ver es que es el **PLN** o procesamiento del lenguaje natural.

Podemos decir que es una rama de la Inteligencia Artificial y la ciencia de la computación que se centra en hacer que las computadoras puedan entender, interpretar y generar lenguaje humano de manera efectiva.

El PLN se refiere a la capacidad de una computadora para entender el lenguaje que hablamos y escribimos, y para realizar tareas como la traducción automática, la generación de texto, la clasificación de texto, el análisis de sentimientos, y otros.

La computadora es capaz de entender el lenguaje que hablamos y escribimos mediante el uso de algoritmos y técnicas de aprendizaje automático que permiten a la misma procesar grandes cantidades de texto y aprender patrones en el lenguaje humano. De esta manera, las computadoras pueden comprender y generar lenguaje humano de una manera eficaz.

Vamos a ver ahora el modelado de lenguaje natural y los LLM. Primeramente podemos manifestar que los modelos de texto se entrenan para predecir la palabra siguiente en una secuencia de texto dada previamente, lo que se conoce como **modelado de lenguaje**, dando como resultado un modelo de lenguaje.

Los LLM son grandes modelos de lenguaje. Viene de *Large Language Model*, que podemos traducir como **modelos de lenguaje grande**. Estos se entrenan con enormes conjuntos de datos de texto y utilizan redes neuronales profundas para procesar la información y aprender patrones en el lenguaje. Por ejemplo, los modelos GPT son LLM, es decir, son grandes modelos de lenguaje.

La tendencia que se va a seguir a partir de ahora no es la de crear modelos de lenguaje cada vez más grandes, porque el problema que tiene esto es que necesita muchísimos recursos; lo que se va a hacer es investigar en nuevos sistemas que permitan, con modelos de lenguaje más pequeños, tener las prestaciones que tienen estos grandes modelos de lenguaje.

El modelado de lenguaje se basa en que las palabras en un idioma no se utilizan de forma aleatoria, sino que las mismas están relacionadas entre sí de una manera predecible, es decir, que en un determinado contexto, después de una determinada palabra, es más probable que venga una palabra que otra.

Por ejemplo, si introducimos este prompt “el cubito de hielo está...” en un modelo de lenguaje, éste lo que va a determinar es que hay más probabilidad de que la siguiente palabra después de “está”, sea la palabra “frío” o la palabra “helado”, y no otras palabras como “caliente” o “blando”. Estas palabras van a tener una probabilidad mucho más baja de ir después de “está” en este contexto.

Uno de los modelos de texto más conocidos son los modelos GPT, que no es ChatGPT, porque ChatGPT es un modelo de diálogo o conversacional que deriva de un modelo de texto, pero es un modelo conversacional.

Bien. Con los conceptos que hemos visto podemos manifestar que los modelos GPT son modelos de lenguaje grande que utilizan redes neuronales tipo transformers para el modelado de lenguaje en tareas específicas de PLN entrenado utilizando grandes cantidades de datos de texto.

Estos modelos son capaces de generar texto coherente y de alta calidad, así como realizar una amplia variedad de tareas de procesamiento del lenguaje natural, como resúmenes, mapas mentales, traducciones, clasificaciones de texto, respuestas a preguntas, etc. Podemos citar como ejemplo el modelo GPT 3, GPT 3.5, GPT 4, y los que vendrán en el futuro. Conoceremos más sobre este tema la clase siguiente.

GPT

Vamos a comenzar diciendo que GPT es un acrónimo de “*Generative Pre-trained Transformer*”, lo que describe su arquitectura y el proceso mediante el que fue entrenado.

Si llevamos al español la traducción literal de GPT, podemos decir que *Transformer* es el nombre de la tecnología; *Pre-trained* quiere decir pre-entrenado; y *generative* significa generativo.

Como dijimos, la T viene de *Transformer*, que es una tecnología que liberó Google de forma Open Source para que cualquier empresa o cualquier desarrollador pudiese utilizarla. La misma lo que hace es procesar el texto de una mejor manera, de una forma mucho más productiva.

Esto ha permitido que cuando entrenamos y aportamos información de diferentes formas a los modelos de Inteligencia Artificial, podemos proporcionar muchísima más información, ya que el modelo es capaz de entenderla mucho mejor. Esto hace que al

final, cuanto más información tenga el modelo, más nutrido esté el mismo, y sea capaz de generar mejores respuestas.

El *pre-trained*, o pre-entrenado, son modelos que ya vienen con un dataset aportado. En el mundo del Machine Learning, que es un subconjunto de la inteligencia artificial que permite a las computadoras aprender de los datos y mejorar con la experiencia sin ser programadas explícitamente, normalmente cuando ustedes generan su propio modelo de Inteligencia Artificial, tienen que aportarle la información; pero en este caso, lo que ya está haciendo Open-AI es generar ese dataset con el que entrenan a GPT. Ese dataset está basado en prácticamente todo internet, o todo lo que han podido conseguir de internet. Entonces ellos hacen esa labor de pre-entrenar el modelo.

Finalmente decimos que es un modelo generativo, ya que implica que ese modelo, cuando genere ese *output*, la respuesta del modelo no va a ser simplemente una respuesta binaria, no va a ser un sí o un no, no va a ser una respuesta numérica, sino que lo que va a generar es un conjunto de letras, de palabras, o de frases, para ser capaces de responder a la petición que ha realizado el usuario.

Aunque no esté definido en el nombre, GPT también es un modelo predictivo, y esta es una de las cosas más importantes.

Lo que hace el modelo, en base a nuestra petición, es generar las siguientes palabras que son más adecuadas para la frase que nosotros estamos aportando, ya sea para continuar la frase, o para generar una respuesta.

En el ejemplo que estamos viendo en pantalla podemos observar que para una frase, el modelo es capaz de ver cuál es la palabra que más posibilidades tiene de ser la palabra correcta. Si esto lo repetimos en una frase entera, al final podremos generar una frase en base a predicciones, y palabra tras palabra, podremos saber si es la palabra correcta.

El modelo lo que hace es predecir, y para predecir lo que tenemos que hacer es aportar mucha información para que esa predicción sea correcta.

Al final, cuanto más información correcta aportemos al sistema, este será mejor a la hora de predecir la información.

Por lo tanto, una de las máximas es la de aportar el suficiente contexto para que el sistema pueda predecir de forma correcta cuál es la respuesta más acertada.

Iremos viendo en el curso que existen diferentes técnicas, y estas son las que debemos utilizar dentro del prompting o el diseño de prompt, para que al final, cuando nosotros nos comuniquemos con la Inteligencia Artificial, esta sea capaz de aportarnos o de contestarnos de la mejor manera.

Ahora nos preguntamos... ¿cómo funciona ChatGPT?

Es importante que entendamos algunos conceptos de por qué ChatGPT está siendo muy popular y esté en este momento de surfear la cresta de la ola de la Inteligencia Artificial.

¿Por qué nosotros consideramos que ChatGPT o GPT es una muy buena opción? GPT desde su inicio ha sido formado mediante una serie de características que han hecho que sin ningún tipo de duda, se convierta en un modelo muy popular en base a la calidad y la facilidad de uso.

Cuando hablamos de calidad, sobre todo hablamos del aprendizaje del modelo, puesto que han sido capaces de entrenar con una cantidad inmensa de información.

Además es un modelo predictivo, es un modelo generador de texto, que es capaz de predecir cuál será la respuesta más adecuada a la pregunta que nosotros estamos haciendo, y por lo tanto, es capaz de generar, en teoría, respuestas de muy alta calidad.

Aparte de ello utilizan tecnologías muy avanzadas, como son los Transformers, que como sabemos, son una nueva forma de utilizar Inteligencia Artificial con el texto; y los prompts, de los que también hablaremos más adelante.

Y finalmente podemos decir que GPT ha triunfado de forma desmesurada por su facilidad de uso. Open-AI ha sido capaz de crear un modelo de Inteligencia Artificial que es capaz de ser utilizado por cualquier persona y no hace falta tener conocimientos técnicos para utilizarlo.

Cuanto más conocimientos tengan, más capaces serán de exprimir este modelo, y eso es lo que vamos a intentar en este curso: dar esas pinceladas, esos conceptos genéricos necesarios, para saber exprimirlo al máximo.

Cuando decimos de datos de entrenamiento, hablamos de qué información tiene ChatGPT, o en este caso GPT en su conocimiento, para poder generar las respuestas adecuadas.

En esta gráfica que tenemos en pantalla podemos observar de una forma muy sencilla cómo el modelo GPT tiene muchísimos más datos de entrenamiento que todos sus predecesores. Estamos viendo simplemente que, a grandes rasgos GPT-3 está entrenado, contiene 175 billones de parámetros.

Pero... ¿qué son los parámetros? sin complicarnos mucho, diremos que son la cantidad de información que tiene el modelo en su base de conocimiento.

Entonces podemos ver, por ejemplo, que versiones anteriores de GPT tienen muchísimos menos parámetros y este es uno de los modelos más grandes que se ha generado.

Cuando hablamos de GPT-3 o GPT-4, nos referimos a diferentes versiones del mismo modelo de Inteligencia Artificial. Es decir, GPT sale al mercado, después existe una versión GPT-2, que es una versión preliminar que Open-AI todavía no había liberado para que la gran cantidad de gente que la está utilizando ahora mismo lo haga de forma sencilla. Y GPT-3 es el primer modelo que hacen público, accesible a todo el mundo, y sobre todo sin esa necesidad de conocimientos técnicos para su uso.

GPT-4 es el siguiente modelo que supera, en este caso, tanto al tamaño de GPT-3 como a la capacidad de generar respuestas de más calidad. No obstante, GPT-4 ahora mismo solo está disponible para usuarios que pagan la suscripción plus de ChatGPT.

Nosotros en este curso sobre todo nos vamos a centrar en la versión gratuita de GPT-3, pero es importante destacar que GPT-4 es un modelo que a día de hoy está disponible, tiene una cantidad de información mayor, y sobre todo es capaz de computarla de mejor manera que GPT-3, que es su predecesor.

La diferencia entre ChatGPT y GPT-3 o GPT-4 es que ChatGPT al final es una interfaz gráfica, es una forma de comunicarnos con estos modelos. Ahora lo vamos a ver más en detalle, pero tenemos que entender que ChatGPT, GPT-3 y GPT-4 son cosas diferentes.

A grandes rasgos, ¿por qué GPT es tan bueno comparado con la competencia? porque es enorme, tiene mucha cantidad de datos, y utiliza la tecnología Transformers, como ya sabemos, que es revolucionaria y que hasta la fecha no se había utilizado como lo hizo Open-AI en este caso.

Si hablamos de tamaño, podemos ver en esta gráfica un poco más segmentada cuáles son los datasets, la cantidad de información que contiene GPT en este caso. Podemos observar que el GPT está formado por un dataset que al final podemos entenderlo como un dataset llamado Common Crawl.

Common Crawl es una base de datos enorme de webs scrapeadas. ¿Saben qué es scrapear? Significa obtener información de páginas web, es accionar y almacenarlo a una base de datos. Entonces estamos ante una base de datos de millones de páginas web almacenadas que posteriormente se le han dado al modelo para que contenga esa información.

En este caso podemos ver como Common Crawl supone el 60% de los datos de entrenamiento de GPT. Y si nos vamos al último registro de esta tabla podemos ver que aparece también Wikipedia, y vemos que la misma supone el 3% de todos los datos de entrenamiento. Es decir, Wikipedia entero supone solo el 3% de toda la información que tiene dentro GPT.

Por lo tanto, podemos hacernos a la idea de lo grande que es este modelo, porque si nosotros pensamos que Wikipedia es enorme y que posee mucha información, pues cuando pensamos que toda esa información solo es representa el 3% de todos los datos que tiene GPT, contiene muchísima información. Entonces Common Crawl al final es casi todo internet, lo que han podido scrapear y exponer para que el propio modelo pueda aprender con una fecha de corte.

ChatGPT tiene una amplia variedad de aplicaciones en diferentes campos:

- Puede ser utilizado para proporcionar soporte al cliente automatizado, respondiendo preguntas frecuentes y ayudando a resolver problemas comunes sin la intervención humana
- Es útil para la generación de textos creativos, como artículos, historias, poemas y más, facilitando a escritores y creadores de contenido el proceso de ideación y redacción
- Puede asistir en la enseñanza proporcionando explicaciones detalladas sobre diversos temas, respondiendo preguntas y ayudando a los estudiantes con sus tareas
- Puede ayudar a los desarrolladores a escribir y depurar código, así como a comprender conceptos complejos de programación
- Puede traducir texto entre diferentes idiomas, lo que lo convierte en una herramienta útil para la comunicación internacional

- Puede analizar grandes volúmenes de texto y extraer información relevante, ayudando a los investigadores y analistas a identificar patrones y obtener conocimientos
- Puede actuar como un asistente virtual, ayudando en la gestión de tareas, recordatorios y proporcionando información relevante a medida que se necesita

En resumen, ChatGPT es una herramienta versátil que aprovecha la Inteligencia Artificial para mejorar la eficiencia y la creatividad en una variedad de tareas relacionadas con el procesamiento del lenguaje natural.

En qué casos ChatGPT funciona de manera correcta

Como aprendimos anteriormente, ChatGPT es un generador de texto predictivo. ¿Qué quiere decir esto? Que está pensado casi en su mayoría para poder comunicarnos con él y respondernos vía texto.

Es un gran comprensor del lenguaje natural, de cómo las personas nos comunicamos en muchísimos idiomas diferentes. Por ende, no es un modelo optimizado para, por ejemplo, procesar grandes cálculos matemáticos. Pero no nos adelantemos y comencemos viendo en qué casos es bueno utilizar dicha herramienta.

Por ejemplo, si queremos generar textos podemos indicarle “escribir post para redes sociales”, “escribir e-mails”, “generar cartas de recomendación”, etc. Es decir, todo lo que sea escribir X, es una muy buena opción.

Si esto lo llevamos a la parte de las páginas web, podemos pedirle, por ejemplo, que escriba fichas de producto, artículos para posicionamiento web, también tema de código de programación, entre otros. O sea, todo lo que tenga que ver con la creación de un texto, por así decirlo.

A su vez es muy bueno generando ideas. Al fin y al cabo, como tiene una base de datos enorme, es capaz de rebuscar en ella y generar nuevas ideas en base a toda esa información. Algo muy parecido a cómo funciona nuestro cerebro, en este caso.

Si nosotros tenemos que crear el asunto de un e-mail para una campaña de e-mail marketing, seguro que pensamos en que ese asunto debe ser llamativo y debe generar un alto CTR para que los usuarios hagan click en ese e-mail, y entonces puedan leer todo nuestro contenido.

Cuando pensamos en esto, automáticamente se nos pueden ocurrir estrategias para llamar la atención, como las típicas frases de “haz clic aquí”, “no te pierdas esta información”, “podrás ver algo de muchísimo valor en muy poco tiempo”, o sea, ganchos que hacen que la gente haga click en el e-mail y ya en el cuerpo del mismo desarrollar un poco más nuestra idea.

Cuando le pedimos al ChatGPT que nos genere asuntos para una campaña de email marketing con una alta capacidad de generar click en base al asunto, lo que va a hacer es ir a su repositorio de información, que como aprendimos, es inmenso, va a buscar todas esas estrategias que se utilizan, y nos propondrá estrategias similares que es capaz de generar a partir de este conocimiento.

Esto es algo muy parecido a lo que nosotros hacemos: en base a nuestra experiencia se nos ocurren nuevos asuntos para los e-mails. En este caso, en base al conocimiento que tiene GPT se le ocurren nuevos asuntos para la campaña de e-mail que le hayamos pedido ayuda.

ChatGPT también es bueno para continuar líneas de pensamiento. Muchos de ustedes seguramente conocerán a los *frameworks* de trabajo, que se encargan de definir la estructura de su futuro proyecto y proporciona las herramientas necesarias que pueden usar como bloques de construcción. Existe, por ejemplo, la metodología Agile, que es una forma de comportarse para intentar desarrollar o trabajar de forma ágil para poder iterar y poder conseguir productos de alta calidad en el mismo tiempo posible.

Por ejemplo, si le establecemos una serie de normas y le decimos “te vas a comportar como una persona que me va a ayudar a hacer Agile”, entonces cada vez que nosotros le preguntemos X cosa, seguirá comportándose como esa especie de maestro, de persona que nos guía dentro de ese marco de pensamiento. Podemos preguntarle “¿podemos hacer esto y seguir trabajando de forma Agile?”. De esta manera, es capaz de continuar esa línea de pensamiento.

Con este ejemplo nos referimos a la parte del diseño de software producto, pero por ejemplo en temas de marketing también hay muchos *frameworks* de trabajo, como los océanos azules para descubrir nuevas oportunidades de marketing. Hay un montón, y lo bueno de ChatGPT es que interactúa muy bien con eso y nos ayuda a desarrollarnos mucho más en ciertas líneas que crean que lo necesitan.

Finalmente podemos decir que ChatGPT es bueno actuando como X profesional. Es decir, le podemos establecer un rol. Por ejemplo, podemos decirle “te vas a comportar

como un experto en marketing offline”, entonces ChatGPT, a partir de la información que nosotros le demos, actuará como ese rol que nosotros hemos establecido.

Como ChatGPT tiene muchos datos, entrenamiento, ha leído un montón de blogs y de artículos de gente especializada un montón de cosas, cuando nosotros le decimos “te vas a comportar como un profesional del marketing offline”, irá a su base de datos, localizará toda la información que tiene sobre ese tipo de profesionales, y entonces será capaz de imitar toda esa información que ha asimilado y de generar datos nuevos, respuestas nuevas, en base a todo eso.

En qué casos ChatGPT no funciona de manera correcta

ChatGPT es un muy buen modelo de generador de texto predictivo; por tanto, podemos asumir que si es muy bueno con texto, es muy bueno con las matemáticas.

Existen muchos modelos de Inteligencia Artificial más buenos con las matemáticas. ChatGPT es el primero que ha destacado en la parte de texto de matemáticas, números, estadística, predicciones, computación avanzada; existen muchísimos desde hace años, y es en lo primero lo que se centraron los investigadores de Inteligencia Artificial.

Es más, es muy habitual que estudiantes de la carrera de matemáticas pasen finalmente a ser ingenieros de Inteligencia Artificial porque tiene aplicaciones muy relacionadas y en la propia carrera de matemáticas se utilizan modelos para hacer grandes cálculos de computación. No obstante, ChatGPT es una muy mala opción para esto, porque realmente está hecho para trabajar con texto, no para hacerlo con matemáticas.

Entonces, muchas veces es capaz de seguir secuencias numéricas o realizar determinadas operaciones sencillas porque es capaz de interpretarlas como si fuese texto, e interpretar los números como texto, y es capaz de inferir ciertos cálculos, pero al final, lo está haciendo un poco de rebote.

Cuando le pedimos cosas un poco más avanzadas, ChatGPT no es una muy buena opción, por lo que todo lo que tenga que ver con pedirle cálculos o conclusiones sobre números, debemos asegurarnos que está realizando una buena deducción.

Esto es muy importante porque muchas veces le pedimos algo y, como nos lo ha dicho la Inteligencia Artificial, tomamos esa información y la asumimos como válida. Pero nosotros mismos tenemos que tener esa capacidad crítica de entender que la respuesta puede ser positiva o no.

Otra cuestión negativa de ChatGPT es que no tiene datos de ejemplo. Es decir, tomando el ejemplo anterior de los e-mails, si nosotros le decimos "dame cinco asuntos para mandar un e-mail", si realmente nosotros no le manifestamos para qué situación queremos conseguir dichos asuntos, cuál es el contexto, quiénes somos, para cuál es la compañía que queremos generarlos, los cinco asuntos que nos va a dar serán totalmente aleatorios y no va a tener una gran capacidad de conseguir nuestro objetivo. Entonces es muy importante que nosotros a ChatGPT le demos un buen contexto.

Finalmente diremos que cuando nosotros le preguntamos a ChatGPT y le hacemos caer en un error con la forma en la que nosotros nos comunicamos, este puede seguir en esa línea de pensamiento y acabar derivando en respuestas totalmente ilógicas.

Por ejemplo, si nosotros le preguntamos si es posible que llueva mañana, nos va a responder que sí, que es posible que llueva mañana. Siempre es posible que llueva el día siguiente, pero eso no significa que vaya a llover. Es decir, si existe un 0,001% de probabilidad, es posible, existen posibilidades, por lo que tenemos que intentar evitar hacer preguntas que nuestra propia pregunta haga que el modelo finalmente termine dando una respuesta totalmente sesgada.

Entonces tenemos que hacer preguntas en las que el modelo sea capaz de razonar cuál es la respuesta. En este caso, en vez de preguntarle si es posible que llueva mañana, deberíamos preguntarle cuál va a ser el tiempo de mañana.

Una cosa muy interesante que sí que se puede hacer es, por ejemplo, buscar en Google un listado del histórico de las lluvias de los últimos 30 años de, por ejemplo, septiembre, y pasarle esa información al modelo. Le ponemos todos los meses y si ha llovido o no ha llovido, y el modelo será capaz de entender esa información y decirnos qué posibilidades existen de que en septiembre llueva en base al histórico.

Errores a evitar al usar ChatGPT

El error número uno es no especificar un rol o especialista para la tarea del comando. Cuando se interactúa con cualquier modelo de lenguaje es muy importante especificar un rol o especialista para la tarea que se ejecuta. Con esto vamos a evitar respuestas genéricas y vamos a tratar de que se centre más en la información que nos va a proporcionar.

Por ejemplo, si quieren crear una campaña de redes sociales, le pueden pedir que su rol sea de un especialista o de un Community manager. Con esto van a evitar muchísimo las respuestas genéricas.

El error número dos sostiene que menos es más. ¿A qué nos referimos con esto? No hacer el prompt demasiado corto, ni demasiado complicado. Un prompt demasiado corto puede dejar a ChatGPT sin la suficiente información para darnos una buena respuesta, pero un comando demasiado largo puede ser demasiado confuso y luego puede hacer que el modelo alucine.

Además, en la versión gratuita hay un límite en la cantidad de tokens, y esto incluye tanto los tokens que utilizamos para la pregunta como para la respuesta; por lo tanto, si es demasiado extenso se van a acabar los tokens de la versión gratuita. Lo ideal es que sean concisos pero claros; traten de ser comprensibles en lo que piden sin agregar detalles innecesarios. Esto ayudará a conservar tokens para una respuesta.

También pueden dividir preguntas complejas. Es decir, si tienen varias preguntas complejas, consideren dividir las en varios prompts más cortos para evitar así los límites de tokens y asegurar respuestas complejas.

El error número tres es no dar un ejemplo claro de cómo quieren los resultados, ya que proporcionar un ejemplo de cómo desean obtener los resultados por parte de ChatGPT puede ser una guía muy buena para el modelo de lenguaje para poderles brindar los resultados tanto en el estilo como en el formato que necesita.

El cuarto error es no dar instrucciones claras de qué hacer. Las instrucciones claras y precisas son fundamentales para que el ChatGPT les pueda dar la respuesta precisa que están necesitando. Sin dichas instrucciones, ChatGPT puede interpretar el prompt de manera errónea, llevando respuestas que no van a cumplir las expectativas de lo que quieren. Podemos citar como ejemplos de instrucciones las siguientes: “analiza el contexto”, “crea 10 títulos de 60 palabras”, “presenta los resultados de una tabla ordenada en columnas”, entre otras.

El error número cinco es pensar que en el primer intento de comando les va a dar la respuesta correcta. Esto es un error muy común que cometemos todos los usuarios cuando le damos el primer comando o prompt a ChatGPT, ya que siempre pensamos que nos va a dar la respuesta perfecta en el primer intento, y esto no es así.

La realidad es que tenemos que hacer muchas pruebas y error hasta obtener el comando perfecto que nos dé la respuesta que estamos buscando. Entonces, es muy importante no pensar que le dieron un prompt y que no les dio la respuesta indicada, o que no funciona el comando, etc. Lo que tienen que hacer es prueba y error hasta perfeccionar este comando y obtener la respuesta que quieren.

El último error es no personalizar ChatGPT. Personalizar esta herramienta para que les proporcione las respuestas en base a su estilo puede mejorar significativamente el tipo de respuestas que les está dando, y esto lo pueden hacer tanto en la versión gratuita como una versión de paga.